

DAFTAR PUSTAKA

- Artificial Intelligence Research Group. (2024). *Roadmap Riset*. Siliwangi University.
- Baly, R., da San Martino, G., Glass, J., & Nakov, P. (2020). We can detect your bias: Predicting the political ideology of news articles. *EMNLP 2020 - 2020 Conference on Empirical Methods in Natural Language Processing, Proceedings of the Conference*. <https://doi.org/10.18653/v1/2020.emnlp-main.404>
- Bird, J. J., Ekárt, A., & Faria, D. R. (2023). Chatbot Interaction with Artificial Intelligence: human data augmentation with T5 and language transformer ensemble for text classification. *Journal of Ambient Intelligence and Humanized Computing*, 14(4). <https://doi.org/10.1007/s12652-021-03439-8>
- Chen, J., Tam, D., Raffel, C., Bansal, M., & Yang, D. (2023). An Empirical Survey of Data Augmentation for Limited Data Learning in NLP. *Transactions of the Association for Computational Linguistics*, 11. https://doi.org/10.1162/tacl_a_00542
- Chen, W. F., Al-Khatib, K., Stein, B., & Wachsmuth, H. (2020). Detecting media bias in news articles using gaussian bias distributions. *Findings of the Association for Computational Linguistics Findings of ACL: EMNLP 2020*. <https://doi.org/10.18653/v1/2020.findings-emnlp.383>
- Chen, W.-F., Al Khatib, K., Wachsmuth, H., & Stein, B. (2020). *Analyzing Political Bias and Unfairness in News Articles at Different Levels of Granularity*. <https://doi.org/10.18653/v1/2020.nlpccs-1.16>
- Feng, S., Park, C. Y., Liu, Y., & Tsvetkov, Y. (2023). From Pretraining Data to Language Models to Downstream Tasks: Tracking the Trails of Political Biases Leading to Unfair NLP Models. *Proceedings of the Annual Meeting of the Association for Computational Linguistics*, 1. <https://doi.org/10.18653/v1/2023.acl-long.656>
- Geng, Y. (2022). Media Bias Detecting based on Word Embedding. *Highlights in Science, Engineering and Technology*, 12. <https://doi.org/10.54097/hset.v12i.1367>
- Ghorbani, A., & Zou, J. (2019). Data shapley: Equitable valuation of data for machine learning. *36th International Conference on Machine Learning, ICML 2019, 2019-June*.

- Gupta, S., Nguyen, H. H., Yamagishi, J., & Echizen, I. (2020). *Viable Threat on News Reading: Generating Biased News Using Natural Language Models*. <https://doi.org/10.18653/v1/2020.nlpccs-1.7>
- Hamborg, F. (2020). Media bias, the social sciences, and NLP: Automating frame analyses to identify bias by word choice and labeling. *Proceedings of the Annual Meeting of the Association for Computational Linguistics*. <https://doi.org/10.18653/v1/2020.acl-srw.12>
- Hamborg, F. (2023). Revealing media bias in news articles: NLP techniques for automated frame analysis. Dalam *Revealing Media Bias in News Articles: NLP Techniques for Automated Frame Analysis*. <https://doi.org/10.1007/978-3-031-17693-7>
- Hamborg, F., Zhukova, A., & Gipp, B. (2019). Automated identification of media bias by word choice and labeling in news articles. *Proceedings of the ACM/IEEE Joint Conference on Digital Libraries, 2019-June*. <https://doi.org/10.1109/JCDL.2019.00036>
- Han, R., Xu, J., & Pan, D. (2022). How Media Exposure, Media Trust, and Media Bias Perception Influence Public Evaluation of COVID-19 Pandemic in International Metropolises. *International Journal of Environmental Research and Public Health*, 19(7). <https://doi.org/10.3390/ijerph19073942>
- Hong, J., Cho, Y., Jung, J., Han, J., & Thorne, J. (2023). Disentangling Structure and Style: Political Bias Detection in News by Inducing Document Hierarchy. *Findings of the Association for Computational Linguistics: EMNLP 2023*. <https://doi.org/10.18653/v1/2023.findings-emnlp.377>
- Horych, T., Wessel, M., Wahle, J. P., Ruas, T., Waßmuth, J., Greiner-Petter, A., Aizawa, A., Gipp, B., & Spinde, T. (2024). *MAGPIE: Multi-Task Media-Bias Analysis Generalization for Pre-Trained Identification of Expressions*. <http://arxiv.org/abs/2403.07910>
- Huang, H., Zhu, H., Liu, W., Gao, H., Jin, H., & Liu, B. (2024). Uncovering the essence of diverse media biases from the semantic embedding space. *Humanities and Social Sciences Communications*, 11(1). <https://doi.org/10.1057/s41599-024-03143-w>
- Jiang, Y., Wang, Y., Song, X., & Maynard, D. (2020). Comparing topic-aware neural networks for bias detection of news. *Frontiers in Artificial Intelligence and Applications*, 325. <https://doi.org/10.3233/FAIA200327>
- Khurana, D., Koli, A., Khatter, K., & Singh, S. (2023). Natural language processing: state of the art, current trends and challenges. *Multimedia*

Tools and Applications, 82(3). <https://doi.org/10.1007/s11042-022-13428-4>

- Krieger, J. D., Spinde, T., Ruas, T., Kulshrestha, J., & Gipp, B. (2022). A domain-Adaptive pre-Training approach for language bias detection in news. *Proceedings of the ACM/IEEE Joint Conference on Digital Libraries*. <https://doi.org/10.1145/3529372.3530932>
- Lee, K., Ippolito, D., Nystrom, A., Zhang, C., Eck, D., Callison-Burch, C., & Carlini, N. (2022). Deduplicating Training Data Makes Language Models Better. *Proceedings of the Annual Meeting of the Association for Computational Linguistics*, 1. <https://doi.org/10.18653/v1/2022.acl-long.577>
- Liu, Y., Ott, M., Goyal, N., Du, J., Joshi, M., Chen, D., Levy, O., Lewis, M., Zettlemoyer, L., & Stoyanov, V. (2019). *RoBERTa: A Robustly Optimized BERT Pretraining Approach*. <http://arxiv.org/abs/1907.11692>
- Menzner, T., & Leidner, J. L. (2024). Experiments in News Bias Detection with Pre-trained Neural Transformers. *Lecture Notes in Computer Science (including subseries Lecture Notes in Artificial Intelligence and Lecture Notes in Bioinformatics)*, 14611 LNCS. https://doi.org/10.1007/978-3-031-56066-8_22
- Murayama, T., Wakamiya, S., & Aramaki, E. (2021). Mitigation of Diachronic Bias in Fake News Detection Dataset. *W-NUT 2021 - 7th Workshop on Noisy User-Generated Text, Proceedings of the Conference*. <https://doi.org/10.18653/v1/2021.wnut-1.21>
- Nadeem, M. U., & Raza, S. (2021). Detecting Bias in News Articles using NLP Models. *Stanford CS224N Natural Language Processing*.
- Otter, D. W., Medina, J. R., & Kalita, J. K. (2021). A Survey of the Usages of Deep Learning for Natural Language Processing. *IEEE Transactions on Neural Networks and Learning Systems*, 32(2). <https://doi.org/10.1109/TNNLS.2020.2979670>
- Pedregosa, F., Varoquaux, G., Gramfort, A., Michel, V., Thirion, B., Grisel, O., Blondel, M., Prettenhofer, P., Weiss, R., Dubourg, V., Vanderplas, J., Passos, A., Cournapeau, D., Brucher, M., Perrot, M., & Duchesnay, É. (2011). Scikit-learn: Machine learning in Python. *Journal of Machine Learning Research*, 12.
- Pembury Smith, M. Q. R., & Ruxton, G. D. (2020). Effective use of the McNemar test. *Behavioral Ecology and Sociobiology*, 74(11). <https://doi.org/10.1007/s00265-020-02916-y>

- Rahali, A., & Akhloufi, M. A. (2023). End-to-End Transformer-Based Models in Textual-Based NLP. Dalam *AI (Switzerland)* (Vol. 4, Nomor 1). <https://doi.org/10.3390/ai4010004>
- Ramadani, Mutiara. S., Khaerudin Kurniawan, & Ahmad Fuadin. (2024). Mengungkap Bias Media dalam Pemberitaan Konflik Israel-Palestina: Sebuah Analisis Konten Kritis. *Jurnal Onoma: Pendidikan, Bahasa, dan Sastra*, 10(1). <https://doi.org/10.30605/onoma.v10i1.3392>
- Raza, S., Reji, D. J., & Ding, C. (2024). Dbias: detecting biases and ensuring fairness in news articles. *International Journal of Data Science and Analytics*, 17(1). <https://doi.org/10.1007/s41060-022-00359-4>
- Reddy, R. R., Duggenpudi, S. R., & Mamidi, R. (2019). *Detecting Political Bias in News Articles Using Headline Attention*. <https://www.ethnologue.com/statistics/size>
- Rodrigo-Ginés, F. J., Carrillo-de-Albornoz, J., & Plaza, L. (2024). A systematic review on media bias detection: What is media bias, how it is expressed, and how to detect it. Dalam *Expert Systems with Applications* (Vol. 237). Elsevier Ltd. <https://doi.org/10.1016/j.eswa.2023.121641>
- Rodrigo-Ginés, F. J., de-Albornoz, J. C., & Plaza, L. (2023). Identifying Media Bias beyond Words: Using Automatic Identification of Persuasive Techniques for Media Bias Detection. *Procesamiento del Lenguaje Natural*, 71. <https://doi.org/10.26342/2023-71-14>
- Saez-Trumper, D., Castillo, C., & Lalmas, M. (2013). Social media news communities: Gatekeeping, coverage, and statement bias. *International Conference on Information and Knowledge Management, Proceedings*. <https://doi.org/10.1145/2505515.2505623>
- Sarker, I. H. (2021). Deep Learning: A Comprehensive Overview on Techniques, Taxonomy, Applications and Research Directions. Dalam *SN Computer Science* (Vol. 2, Nomor 6). <https://doi.org/10.1007/s42979-021-00815-1>
- Shannon, C. E. (1948). A Mathematical Theory of Communication. *Bell System Technical Journal*, 27(3). <https://doi.org/10.1002/j.1538-7305.1948.tb01338.x>
- Shri Bharathi, S. V., & Geetha, A. (2019). Determination of news biasedness using content sentiment analysis algorithm. *Indonesian Journal of Electrical Engineering and Computer Science*, 16(2). <https://doi.org/10.11591/ijeecs.v16.i2.pp882-889>

- Siino, M., Tinnirello, I., & La Cascia, M. (2024). Is text preprocessing still worth the time? A comparative survey on the influence of popular preprocessing methods on Transformers and traditional classifiers. *Information Systems*, 121. <https://doi.org/10.1016/j.is.2023.102342>
- Sokolova, M. (2009). ScienceDirect - Information Processing & Management: A systematic analysis of performance measures for classification tasks. *Information Processing & Management*.
- Spinde, T., Krieger, J. D., Ruas, T., Mitrović, J., Götz-Hahn, F., Aizawa, A., & Gipp, B. (2022). Exploiting Transformer-Based Multitask Learning for the Detection of Media Bias in News Articles. *Lecture Notes in Computer Science (including subseries Lecture Notes in Artificial Intelligence and Lecture Notes in Bioinformatics)*, 13192 LNCS. https://doi.org/10.1007/978-3-030-96957-8_20
- Spinde, T., Plank, M., Krieger, J. D., Ruas, T., Gipp, B., & Aizawa, A. (2021). Neural Media Bias Detection Using Distant Supervision with BABE - Bias Annotations by Experts. *Findings of the Association for Computational Linguistics, Findings of ACL: EMNLP 2021*. <https://doi.org/10.18653/v1/2021.findings-emnlp.101>
- Stafford, T. (2014). *Psychology: Why Bad News Dominates the Headlines*. 29 July.
- Zhou, J., & Bhat, S. (2021). Paraphrase Generation: A Survey of the State of the Art. *EMNLP 2021 - 2021 Conference on Empirical Methods in Natural Language Processing, Proceedings*. <https://doi.org/10.18653/v1/2021.emnlp-main.414>
- Zhu, Y., Sheng, Q., Cao, J., Li, S., Wang, D., & Zhuang, F. (2022). Generalizing to the Future: Mitigating Entity Bias in Fake News Detection. *SIGIR 2022 - Proceedings of the 45th International ACM SIGIR Conference on Research and Development in Information Retrieval*. <https://doi.org/10.1145/3477495.3531816>